



Subject: Internship Proposal

<i>ID</i>	PTI_EN_Finocchio Giovanni_07/08/2026 10.58.42
<i>Data</i>	07/08/2026 10.58.42

Project Supervisor

<i>Surname</i>	Finocchio
<i>Name</i>	Giovanni
<i>Department</i>	MIFT
<i>Laboratory</i>	Magnetism and Spintronics
<i>E-mail</i>	gfinocchio@unime.it
<i>Phone number</i>	3286264205

Project Co-Supervisor

<i>Surname</i>	
<i>Name</i>	
<i>Job Position</i>	
<i>Department</i>	

<i>Laboratory</i>	
<i>E-mail</i>	
<i>Phone number</i>	

Project details

<i>Title</i>	Benchmarking the Learning Efficiency of Sequential Neural Models and Reservoir Computing
<i>Detailed description:</i> It is common to evaluate a machine learning model by measuring its predictive accuracy on a test dataset. This approach favours complex models that generalise well, but it rarely accounts for the speed and data efficiency of the learning process, which are essential components of intelligence and are especially important for low-data applications. This project studies and experimentally compares established sequential supervised models, such as Recurrent Neural Networks (RNNs), Long Short-Term Memory networks (LSTMs), Gated Recurrent Units (GRUs) and Transformers, with less common alternatives based on reservoir computing, namely Echo State Networks (ESNs) and Reservoir Cellular Automata (ReCA). The work relies on a recently proposed benchmark of progressively more difficult sequential tasks, together with a data-efficiency metric, the Weighted Average Data Efficiency (WADE), which measures how quickly a model reaches a set of test-accuracy checkpoints during training. The internship consisted in studying these models, reproducing the published benchmark results, and analysing and comparing the models in terms of both accuracy and learning efficiency. Internship Objectives The main objective of this internship is to acquire a solid understanding of sequential neural models and reservoir computing, and to reproduce and analyse a learning-efficiency benchmark. The work may include one or more of the following activities: (i) Literature review of sequential supervised models (RNN, LSTM, GRU, Transformer) and of reservoir computing approaches (Echo State Networks and Reservoir Cellular Automata), covering the concepts of hidden state, gating, attention, and frozen reservoirs with a trained read-out layer. (ii) Study of the Weighted Average Data Efficiency (WADE) metric and of its role in measuring the learning speed of a model rather than its final accuracy.	



- (iii) Set-up of the experimental environment in Python (PyTorch, Poetry) and reproduction of the benchmark code, tables and figures from the reference work.
- (iv) Description and analysis of the benchmark tasks, ranging from simple periodic-pattern recognition to counting and question-answering tasks, together with the IMDB text-classification task.
- (v) Experimental evaluation and comparison of the studied models in terms of both accuracy and learning efficiency, and interpretation of the differences between fully-supervised and reservoir-based models.
- (vi) Preparation of visual summaries, tables and concise technical documentation of the obtained results.

<i>Duration (month – max 12)</i>	4
<i>Duration (hours)</i>	150
<i>Open positions</i>	1

Internship Skills

Technical requirements: basic knowledge of machine learning and neural networks; programming in Python

<i>Other skills</i>	PyTorch, Jupyter
---------------------	------------------